

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Pragmatic Explanation of Causal Inference from Correlational Statements

Permalink

<https://escholarship.org/uc/item/6089f9kx>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Zur, Amir

Goodwin, Emily

Shain, Cory

et al.

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Pragmatic Explanation of Causal Inference from Correlational Statements

Anonymous CogSci submission

Abstract

People sometimes infer causal relationships from correlational evidence. For example, upon hearing *cancer is associated with smoking*, listeners might conclude smoking causes cancer. Without background information to inform priors, readers tend to interpret subjects of *associated with* as effects, rather than causes (Gershman & Ullman, 2023). We replicate this preference across multiple correlational predicates, and then derive it within a rational speech acts model of pragmatic inference. Central to our analysis is the observation that English asymmetrically constrains speakers: a speaker who plans to utter an effect in subject position has fewer, higher-cost alternatives compared to a speaker who plans to utter a cause in subject position. This asymmetry drives the pragmatic listener to infer that the use of a symmetric predicate signals that the subject is an effect. Our analysis additionally accounts for some of the variability in preferences for the subject-as-effect interpretation across different correlational predicates.

Keywords: pragmatic inference; RSA; causal language

Introduction

Although it is commonly observed that correlation does not imply causation, people do use correlational statements to inform their causal models of the world. That is, in a conversation about topic *A*, the statement “*A is associated with B*” can lead a listener to interpret that a causal relationship holds between *A* and *B*. As Lassiter and Franke (2024) observe, this inference is plausibly guided by pragmatic assumptions: a speaker who endeavors to say relevant things would not mention a correlation they believe is spurious (except in a conversation about surprising, accidental correlations). Because non-causal correlations are not informative, a listener can infer that statements about correlations are intended by the speaker as relevant to a causal world model.

Gershman and Ullman (2023) report that English speakers reading about associations between nonce words (*A is associated with B*) preferred to interpret the subject as an effect rather than cause. Why participants have this preference is unclear, given that the predicate describes a symmetric relationship. That is, causes and effects are equally plausible in subject position when readers have clear expectations about the likely direction of causation (*smoking is associated with cancer* and *cancer is associated with smoking*).

Lassiter and Franke (2024) argue that the subject-as-effect preference (in the absence of such prior expectations) follows from how readers reason about missing discourse context.

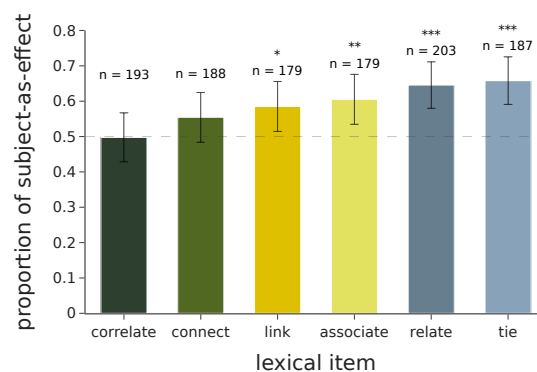


Figure 1: The proportion of participants who selected the subject-as-effect interpretation in Experiment 1. Asterisks indicate statistically significant differences from chance (.5) in a 2-tailed *t*-test: * $p < 0.05$, ** $p < 0.01$, *** $p < 0.0001$.

Their argument assumes two things: first, that readers will infer the topic *t* of the conversation is the subject of the sentence; and second, that readers will infer the question under discussion (QUD) is “what causes *t*?”. The first assumption is motivated by canonical English information structure. But the second assumption does not have an independent motivation, and an attempted topicality manipulation failed to reverse the effect (Experiment 1b, Lassiter & Franke, 2024, p.3-4). Moreover, prior work on the subject-as-effect preference focused only on the term *associated*, leaving open the generality of this effect across other correlational predicates.

In this paper, we advocate and test an alternative explanation of the subject-as-effect preference, drawing on the rational speech acts (RSA) framework (Frank & Goodman, 2012). The central insight behind our explanation is that, of the English predicates which express causation, the majority take a subject which is a cause (e.g. *causes*, *increases*). These predicates are collectively more frequent and syntactically simpler than those which take an effect as subject (e.g. *relies on*, *is predicted on*). Thus, when the topic of discussion is an effect, the speaker must choose between producing a higher-cost predicate or putting the topic in non-subject position, violating typical English information structure. Listeners who model

speakers' utterance alternatives are more likely to infer that the subject of a symmetric predicate is an effect, reasoning that if it was a cause, the speaker would have had more convenient alternatives.

We present two sets of experimental results. First, we replicated the subject-as-effect interpretation preference in five additional correlation predicates, ruling out the possibility that the subject-as-effect interpretation is an idiosyncratic, lexically-specified preference of *associated with*. An RSA model of our alternatives-based account correctly predicted the observed preferences in these experiments. In addition to the subject-as-effect preference observed by Lassiter and Franke (2024), the RSA model explained some of the variability between different correlational predicates (reported in Experiment 1 and visualized in Figure 1). We also tested an additional QUD manipulation, which Lassiter and Franke (2024)'s account predicts would reverse the subject-as-effect preference; it did not reverse the preference, providing no evidence for QUD-specific predictions.

In summary, listeners infer causation from correlational statements (subject-as-effect), even when they lack priors over causal world models, by reasoning rationally about the speaker's intention given the information structure of English. A pragmatic listener prefers the subject-as-effect interpretation due to the number of subject-as-cause alternatives available to the speaker. This pragmatic inference accounts for the subject-as-effect preference across QUD manipulations and different correlational predicates.

An Alternatives-Based Account

We propose that readers arrive at the subject-as-effect interpretation through Gricean reasoning about alternative statements that the speaker could have uttered (Gotzner & Romoli, 2022; Grice, 1975). Our account rests on two observations. First, speakers typically place the topic of discussion in subject position, a well-documented feature of English information structure (Lambrecht, 1994). Second, English constrains how speakers can express causal relationships: of the predicates that express a causal relation, the vast majority take a subject which is a cause rather than effect (e.g., *A causes B*, *A produces B*). Predicates that take an effect as subject are rarer and often have more words (e.g., *A results from B*, *A is predicated on B*).

Together, these facts create an asymmetry. When the topic of discussion is a cause, speakers have many low-cost ways to express a causal relation putting the topic in subject position. When the topic is an effect, however, speakers are forced to either use a costlier predicate or violate typical information structure. A pragmatic listener hearing a correlational predicate can invert this reasoning. Had the speaker wanted to convey that the subject is the cause, they would've had many convenient options available, yet they used none of them. This is more likely if the speaker intended the subject to be the effect, for which fewer low-cost alternatives exist.

	subject as cause	subject as effect
number of verbs	71 (87%)	11 (13%)
token frequency	391 (60%)	258 (40%)
number of unigrams	53 : 18 (75%)	4 : 7 (36%)

Table 1: In the BECAUSE dataset, predicates requiring causes in subject position outnumber those that require effects in subject position; this set collectively has a higher token frequency and unigram ratio (Dunietz et al., 2017).

Evidence of asymmetry in English causal predicates We investigated causal predicates in the Bank of Effects and Causes Stated Explicitly (BECAUSE), an annotated dataset of lexical markers for causal constructions (Dunietz et al., 2017). Table 1 summarizes our findings. Predicates which require causes in subject position had a higher type count than those which require effects (71 vs 11 lemmas), a higher total token count (60% vs. 40%), and were more often syntactically simple (75% vs. 36% within each respective category are unigrams). By all of these standard measures of accessibility, we observe that utterances with effects in subject position are costlier to produce. We do not model why speakers might opt for a correlational predicate to express causality. It may communicate uncertainty over causal strength, or express evidential basis (Lassiter & Franke, 2024). Our account addresses a narrower question: *given* that a speaker has used a correlational predicate, how do listeners infer causal direction?

Accounting for predicate-specific variation Our account predicts that the strength of the inference that *B* causes *A* should depend on which alternatives are salient to the listener. This parallels findings with scalar implicatures, where different lexical scales yield strikingly different rates of pragmatic inference (Gotzner et al., 2018; Van Tiel et al., 2016), and RSA models have successfully captured this variation by manipulating the alternative set (Franke, 2014; Peloquin & Frank, 2016). Just as the strength of scalar implicatures varies with properties of the lexical scale, we predict that the strength of the subject-as-effect inference will vary across correlational predicates. Specifically, we propose that the *active* form of the predicate enters the listener's consideration set to the extent that it is felicitous for expressing collinearity. Because active correlational forms are also symmetric, their presence dilutes the asymmetric signal from causal predicates. This predicts variation across predicates. For the predicate *associate*, the active form *A associates with B* does not felicitously express collinearity; it instead conveys social affiliation. So listeners reason primarily about causal alternatives, producing a stronger subject-as-effect preference. For *correlate*, the active form *A correlates with B* naturally expresses collinearity, providing a salient symmetric alternative that weakens the subject-as-effect inference.

RSA model We formalize our alternatives-based account within the Rational Speech Acts (RSA) framework, which models pragmatic inference as recursive reasoning between speakers and listeners (Degen, 2023; Frank & Goodman, 2012). The model has three components: a literal listener that interprets utterances against their truth-conditional content, a pragmatic speaker that chooses utterances based on informativity and cost, and a pragmatic listener that inverts the speaker’s reasoning.

Literal listener Following standard practice in probabilistic models of causal inference (Griffiths & Tenenbaum, 2005, e.g.), we restrict the listener’s hypothesis space to the minimal set of causal graphs over the variables the utterance introduces. For an utterance about A and B , this yields three graphs: $G_{A \perp B}$, $G_{A \rightarrow B}$, and $G_{B \rightarrow A}$. Common-cause and common-effect structures require additional latent variables not introduced by the utterance and so fall outside the partition. The literal listener updates a uniform prior over G by the truth-conditional content of the utterance,

$$P_{L_0}(G | u) \propto \llbracket u \rrbracket(G). \quad (1)$$

A causal predicate such as A causes B is true only in $G_{A \rightarrow B}$ and so picks out a single world. A correlational predicate such as A is associated with B asserts only that A and B are not independent; its truth conditions are silent about direction. It rules out $G_{A \perp B}$ but distributes posterior mass equally over the two directional worlds:

$$P_{L_0}(G_{A \rightarrow B} | u_{\text{corr}}) = P_{L_0}(G_{B \rightarrow A} | u_{\text{corr}}) = \frac{1}{2}. \quad (2)$$

Crucially, the literal listener has no preference between the two causal directions; the subject-as-effect preference must arise from pragmatic reasoning at higher levels of the model.

Pragmatic speaker To model the pragmatic speaker, we consider a distribution that is proportional to the literal listener. Although the speaker might intend to communicate more than cause and effect (e.g., the evidential basis for the causal relation), we treat these as separate variables that aren’t relevant to the inference of causal direction. The probability that the speaker chooses utterance u to convey the causal direction c is then

$$P_S(u | c) \propto \exp[\alpha \log P_{L_0}(c | u) - \beta C(u) + \gamma F(u)] \quad (3)$$

where α , β , and γ are scalars, $C(u)$ is the cost of the utterance, and $F(u)$ is the felicity of the utterance for expressing collinearity. When the speaker wishes to communicate that the subject is the effect, all options are relatively high cost (see below) and thus the relative cost of producing a correlational statement (instead of a causal one) is reduced, leading to a higher production probability.

The cost of an utterance is the inverse of its frequency as in Table 1.

$$C(u) = \frac{1}{n_u} \quad (4)$$

where n_u is the number of lemmas in the BECAUSE dataset that have the corresponding construal (first row in Table 1). For correlational predicates, we fix $C(u) = 0.5$.

We compute the felicity of an utterance as follows.

$$F(u) = f_p \cdot \mathbb{1}[u = \text{active form of } p] \quad (5)$$

where f_p depends on the predicate p heard by the listener. A higher f_p value means that the statement is felicitous in active form for expressing collinearity (e.g., A correlates with B). To estimate f_p , we collected human judgments of felicity; we report on these experiments in the RSA Analysis section.

Pragmatic listener The final stage of our RSA model, the pragmatic listener, inverts the reasoning process of the speaker. Formally,

$$P_L(c | u) \propto P_S(u | c) \quad (6)$$

When hearing A is associated with B , a pragmatic listener would reason that the speaker had lower-cost alternatives to communicate $A \rightarrow B$ (e.g., A causes B) than to communicate $B \rightarrow A$ (e.g., B is predicated on A). Hence, the pragmatic listener would assign greater probability to A being the effect over A being the cause.

In the next section, we show that the subject-as-effect preference first reported by Gershman and Ullman (2023) generalizes robustly to other correlational predicates.

Experiment 1

Because our explanation for the subject-as-effect preference applies across correlational predicates, we tested this generality by replicating the experimental setup of Gershman and Ullman (2023) with different correlational predicates.

Methods

We recruited 1,200 participants on Prolific, who each completed one experimental trial. Participants read “Suppose you read the following piece of information: ‘ A is predicate B ’. Which of the following is more likely?” and chose between randomly-ordered choices “ A causes B ” and “ B causes A ”. Noun pairs A and B were randomly sampled from those used by Gershman and Ullman (2023) and Lassiter and Franke (2024). The predicate was sampled from the following: *associated with*, *correlated with*, *linked to*, *tied to*, *related to*, *connected to*.

To familiarize participants with the format of the trials, participants first completed a practice trial with an unambiguous question: “Suppose you read the following piece of information: ‘While Maria was typing her report, the phone rang’” and choosing between “Maria began her report first” and “The phone rang first”.

Following Gershman and Ullman (2023), we included an exit survey where participants were asked to describe in their own words what they had done during the experiment; we excluded participants who gave nonsensical answers (“Compete a sentences”) and participants who indicated that their native

language was not English. All experiments in this paper were coded using JsPsych (De Leeuw et al., 2023), and used the DataPipe plugin for data collection¹.

Results and Discussion

After exclusions, 1,135 participants remained. For predicates *correlated with* and *connected to*, the rate of subject-as-effect interpretation was not significantly different from chance; all other predicates led to a significant preference for subject-as-effect (Figure 1).

The participants in our study showed a preference for the subject-as-effect interpretation across different correlational predicates. Hence, the subject-as-effect preference is not restricted to the verb *associate*. Two exceptions are *connect* and *correlate*, which were not significantly different from chance. We discuss two hypotheses for the symmetric causal direction implied by *correlate*, and show how the felicity parameter of our RSA model interacts with them.

The first hypothesis is that the subject of *correlate* is not constrained in the same way as other correlational predicates. That is, *correlate* can be used in both an active and a passive construction: *A correlates with B* might convey a causal relation, whereas *A associates with B* would not (*associates with* is grammatical but does not express collinearity). Hence, the speaker would have equally good alternatives for expressing $A \rightarrow B$ and $B \rightarrow A$, and a pragmatic listener would infer that these causal directions are equally likely.

A large felicity parameter f_p for $p = \textit{correlates}$ accounts for this hypothesis: since *correlates with* is felicitous to a discussion about collinearity, the listener considers it a more likely alternative than causal predicates. Meanwhile, *associates with* has a low felicity parameter, helping explain the difference in the subject-as-effect preference between the two predicates.

The second hypothesis concerning *correlates* is that, indeed, correlation does not imply causation – or rather, the statement *A is correlated with B* is taken to imply that there is no direct causal relation between *A* and *B*. In this case, listeners would not infer a causal relation from *correlate*, and so their choice of causal direction would be arbitrary. Independent evidence supports this interpretation of *correlate*: in a corpus study of medical literature, Haber et al. (2022) found that *correlate* is most often used to convey no causal relation between its arguments.

The felicity parameter similarly accounts for the second hypothesis. A large value for f_p would mean that pragmatic listeners do not consider causal predicates as alternatives for $p = \textit{correlate}$. Since correlational predicates are symmetric, the RSA model would predict no subject-as-effect preference for the statement *A is correlated with B*.

In summary, participants in our study showed a significant subject-as-effect preference across lexical items, with the exception of *correlate* and *connect* (see Figure 1). The participant preferences were qualitatively consistent with an

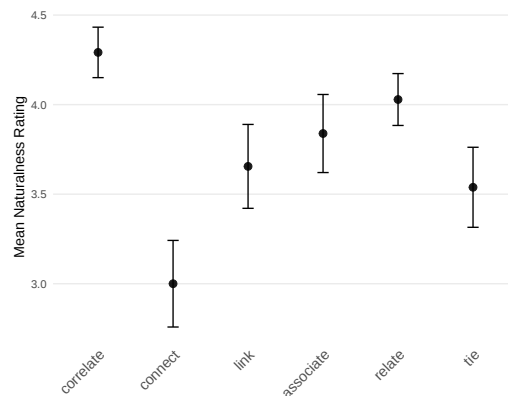


Figure 2: Mean naturalness judgments for each of the predicates in active voice.

alternatives-based account. In the next section, we model these data quantitatively using predictions from the RSA model.

RSA Analysis

To provide a computational explanation of the subject-as-effect preference, we fit the RSA model in Equation 6 to the human preferences we collected in Experiment 1. We measured how closely the model accounted for the graded variability in the subject-as-effect interpretation across correlational predicates.

Methods

The free parameters of our RSA model are α , which controls the speaker informativity, β , which controls the importance of ease of production to the speaker, and γ , which controls the active-form felicity for each correlational predicate. We collected human felicity judgments to estimate the felicity for each predicate, f_p , and fit the α and β parameters to the participant responses in Experiment 1 using the memo library (Chandra et al., 2025).

To measure the felicity of active forms, we collected naturalness judgments from an additional 180 participants on Prolific. Each participant was presented with a slider, and asked “You have learned that two things, *A* and *B*, are **collinear**. This means that they tend to increase and decrease together. How natural is the following as a description of this relationship? *A predicate B*”. Nouns *A* and *B* were taken from the same set as the other experiments; predicates were one of: *correlates with*, *connects to*, *links to*, *associates with*, *relates to* or *ties to*. Slider edges were labeled with “Very Unnatural” and “Perfectly Natural”, and responses were coded automatically by the plugin as whole values from one to five.

After excluding three participants who did not report being native speakers of English on the exit survey, we were left with 177 participants. Figure 2 reports the felicity scores.

We computed f values by taking the average participant felicity judgment for each predicate, and applying the softmax

¹<https://pipe.jspsych.org>

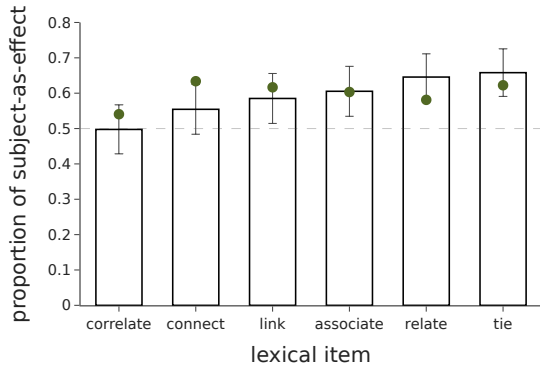


Figure 3: Experiment 1 results (outlined bars), and corresponding RSA predictions (green circles).

function across these values. Then, we fit the RSA model to the average participant response for each predicate. To predict participant responses, we optimized the α and β parameters with gradient descent, minimizing root mean-squared error (RMSE) between the RSA prediction and the rate of subject-as-effect interpretation for each predicate.

Results and Discussion

After fitting to participant responses, the RSA model achieved a root mean-squared error (RMSE) of 0.0510 with optimal parameters $\alpha \approx 1.72$, $\beta \approx 7.53$, and $\gamma \approx 19.50$. Figure 3 shows the participant choices, along with predictions from the RSA model. The RSA predictions fell within the 95% confidence intervals for all lexical items.

The RSA model correctly predicted a subject-as-effect preference. Moreover, the predictions varied by correlational predicate, with the lowest subject-as-effect preference for *correlate*.

These results support the hypothesis that listeners reason about alternative statements when inferring causation from correlational statements. In the next section, we distinguish between the QUD and the alternatives-based accounts of the subject-as-effect preference through an additional QUD manipulation.

Experiment 2

We have demonstrated that the subject-as-effect interpretation is not an idiosyncratic property of *associated*, which is consistent with pragmatic accounts. Recall that, under Lassiter and Franke’s (2024) account, readers preferentially infer a QUD of “what causes *topic*?”, leading to the subject-as-effect preference. However, manipulating the topicality directly did not reverse the preference, and it remains unclear why readers would prefer QUDs about causes rather than effects. In this section, we investigate a possible explanation for the QUD

preference: it may have been induced by the alternatives presented to participants.

Recall from Experiment 1 that participants chose between options phrased with *cause*: “A causes B” or “B causes A”. This particular choice of predicate might privilege QUDs about causes. We theorize that participants asked to choose between “A is an effect of B” and “B is an effect of A”, might be more likely to infer a QUD of “What is an effect of *topic*”? than “What is a cause of *topic*?”. By varying the choices presented to participants, we manipulated the QUD while keeping the same set of possible alternative utterances. Hence, the QUD account predicts that participant preferences would flip from subject-as-effect to subject-as-cause, whereas the alternatives-based account predicts no change in participant responses.

Methods

We recruited an additional 600 participants from prolific. Following Experiment 1, participants read “Suppose you read the following piece of information: ‘A is *predicate* B’. Which of the following is more likely?”. Half the participants then chose between “A causes B” and “B causes A”; the other half chose between “A is an effect of B” and “B is an effect of A”. After making their selection, an additional slider appeared where participants could answer “How certain are you?” Participants also completed a practice trial and exit survey, as in Experiment 1.

Results and Discussion

Eight participants’ data were lost, due to errors with the data saving pipeline. After excluding an additional 15, we were left with 577 participants. The proportion of subject-as-effect interpretations is shown in Figure 4. The phrasing of the options presented to participants had no significant effect on subject-as-effect preference ($p = 0.21$).

The results of this experiment are consistent with a model where participants’ subject-as-effect preference is not due to the QUD they infer. In particular, for the QUD approach to explain the subject-as-effect preference in our experiments, participants must recover by recovering a QUD of “what is the cause of *topic*?” even when faced with alternatives phrased in terms of effects rather than causes.

We note, however, that QUD and topicality may simply be too difficult to manipulate in this experiment paradigm. That is, it is very difficult to construct a discourse context that makes one noun phrase unambiguously and intuitively the discourse topic, without communicating any information that might inform whether it is more likely to be a cause or an effect. In fact, Lassiter and Franke (2024) tested such a manipulation, finding that even small amounts of background knowledge (such as the suggestion that one of the arguments was a drug or a disease) were sufficient to remove the subject-as-effect preference. The difficulty of experimentally validating what QUD participants converge on makes it hard to conclusively rule out the role of QUD in inference. Instead, we see our

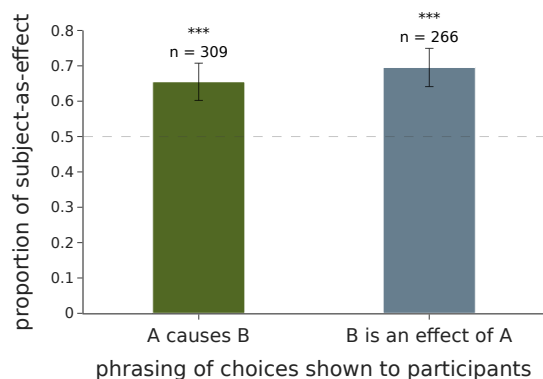


Figure 4: The proportion of participants who selected the subject-as-effect interpretation in Experiment 2. Asterisks indicate statistically significant differences from chance (.5) in a 2-tailed t -test: * $p < 0.05$, ** $p < 0.01$, *** $p < 0.0001$.

experiment as an additional piece of evidence, in keeping with the previously failed QUD manipulations.

Discussion

This paper investigates the tendency of readers to interpret subjects of *associated with* as effects, rather than causes. This preference has been explained as an effect of how people reason about what a likely QUD would be, without having discourse context. In our first experiment, we replicated the effect for predicates other than *associated with*, broadening the empirical picture to include the class of predicates conveying correlation as a whole. This supports pragmatic accounts of the effect, ruling out the concern that the effect is merely an idiosyncratic property of one lexical item. We also discovered significant variation in the effect across such predicates, which is not predicted by the QUD-based account.

We then introduce an RSA model, grounded in the assumption that listeners reason recursively about the utterance alternatives available to a speaker. Using this model, we can account for both the original subject-as-effect preference displayed by *associated with*, as well as some of the variance among predicates, and remove the need to assume that readers preferentially infer a particular QUD.

Although our RSA model accounts for why listeners might infer causation from correlational statements, we do not model why speakers would opt to use correlational statements to express causation – even though correlational statements are often employed for causal associations (Adams et al., 2017; Haber et al., 2022). One hypothesis is that speakers use correlational statements as evidential markers. For instance, they might imply the degree of the speaker’s commitment to a causal relation or for the evidential basis of one. An explanation of the production of correlational statements would need

to take the speaker’s epistemic uncertainty about the relation between A and B into account.

Conclusion

People use correlational statements to inform their causal models of the world, and rely on priors to identify cause and effect. However, absent these priors, listeners are more likely to interpret the subject as the effect rather than the cause. We propose that this subject-as-effect preference is due to an asymmetry in English: causal predicates are more likely to assign cause than effect to their subject argument.

Grounding the alternatives-based account in the RSA framework, we demonstrated how pragmatic listeners reason about speaker alternatives to infer the subject-as-effect interpretation. We reported participant choices across correlational predicates and QUD manipulations, which establish the subject-as-effect preference as a consistent pragmatic inference.

Our work is part of a broad literature on how language shapes the way we express and infer causality, and computational models of the pragmatic reasoning processes that come about from these constraints (Beller & Gerstenberg, 2025; Beller et al., 2020; Gerstenberg et al., 2021; Talmy, 1988; Wolff, 2007). To study how listeners infer causation from correlation, we proposed and tested an account of how the asymmetry of causal predicates in English leads listeners to interpret the subject of the sentence as the effect modulo additional context.

References

- Adams, R. C., Sumner, P., Vivian-Griffiths, S., Barrington, A., Williams, A., Boivin, J., Chambers, C. D., & Bott, L. (2017). How readers understand causal and correlational expressions used in news headlines. *Journal of Experimental Psychology: Applied*, 23(1), 1–14. <https://doi.org/10.1037/xap0000100>
- Beller, A., Bennett, E., & Gerstenberg, T. (2020). The language of causation. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 42.
- Beller, A., & Gerstenberg, T. (2025). Causation, meaning, and communication. *Psychological Review*.
- Chandra, K., Chen, T., Tenenbaum, J. B., & Ragan-Kelley, J. (2025). A domain-specific probabilistic programming language for reasoning about reasoning (or: A memo on memo). *Proc. ACM Program. Lang.*, 9(OOPSLA2). <https://doi.org/10.1145/3763078>
- De Leeuw, J. R., Gilbert, R. A., & Luchterhandt, B. (2023). jsPsych: Enabling an Open-Source Collaborative Ecosystem of Behavioral Experiments. *Journal of Open Source Software*, 8(85), 5351. <https://doi.org/10.21105/joss.05351>
- Degen, J. (2023). The rational speech act framework. *Annual Review of Linguistics*, 9(1), 519–540.

- Dunietz, J., Levin, L., & Carbonell, J. G. (2017). The because corpus 2.0: Annotating causality and overlapping relations. *Proceedings of the 11th linguistic annotation workshop*, 95–104.
- Frank, M. C., & Goodman, N. D. (2012). Predicting pragmatic reasoning in language games. *Science*, 336(6084), 998–998.
- Franke, M. (2014). Typical use of quantifiers: A probabilistic speaker model. *Proceedings of the annual meeting of the cognitive science society*, 36(36).
- Gershman, S. J., & Ullman, T. D. (2023). Causal implicatures from correlational statements. *PloS one*, 18(5), e0286067.
- Gerstenberg, T., Goodman, N. D., Lagnado, D. A., & Tenenbaum, J. B. (2021). A counterfactual simulation model of causal judgments for physical events. *Psychological review*, 128(5), 936.
- Gotzner, N., & Romoli, J. (2022). Meaning and alternatives. *Annual Review of Linguistics*, 8, 213–234.
- Gotzner, N., Solt, S., & Benz, A. (2018). Scalar diversity, negative strengthening, and adjectival semantics. *Frontiers in psychology*, 9, 1659.
- Grice, H. P. (1975). Logic and conversation. In *Speech acts* (pp. 41–58). Brill.
- Griffiths, T. L., & Tenenbaum, J. B. (2005). Structure and strength in causal induction. *Cognitive Psychology*, 51(4), 334–384. <https://doi.org/https://doi.org/10.1016/j.cogpsych.2005.05.004>
- Haber, N. A., Wieten, S. E., Rohrer, J. M., Arah, O. A., Tennant, P. W. G., Stuart, E. A., Murray, E. J., Pilleron, S., Lam, S. T., Riederer, E., Howcutt, S. J., Simmons, A. E., Leyrat, C., Schoenegger, P., Booman, A., Dufour, M.-S. K., O'Donoghue, A. L., Baglini, R., Do, S., . . . Fox, M. P. (2022). Causal and Associational Language in Observational Health Research: A Systematic Evaluation. *American Journal of Epidemiology*, 191(12), 2084–2097. <https://doi.org/10.1093/aje/kwac137>
- Lambrecht, K. (1994). *Information structure and sentence form: Topic, focus, and the mental representations of discourse referents* (Vol. 71). Cambridge university press.
- Lassiter, D., & Franke, M. (2024). The rationality of inferring causation from correlational language. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 46.
- Pelouquin, B., & Frank, M. C. (2016). Determining the alternatives for scalar implicature. *Proceedings of the annual meeting of the cognitive science society*, 38.
- Talmy, L. (1988). Force dynamics in language and cognition. *Cognitive science*, 12(1), 49–100.
- Van Tiel, B., Van Miltenburg, E., Zevakhina, N., & Geurts, B. (2016). Scalar diversity. *Journal of semantics*, 33(1), 137–175.
- Wolff, P. (2007). Representing causation. *Journal of experimental psychology: General*, 136(1), 82.